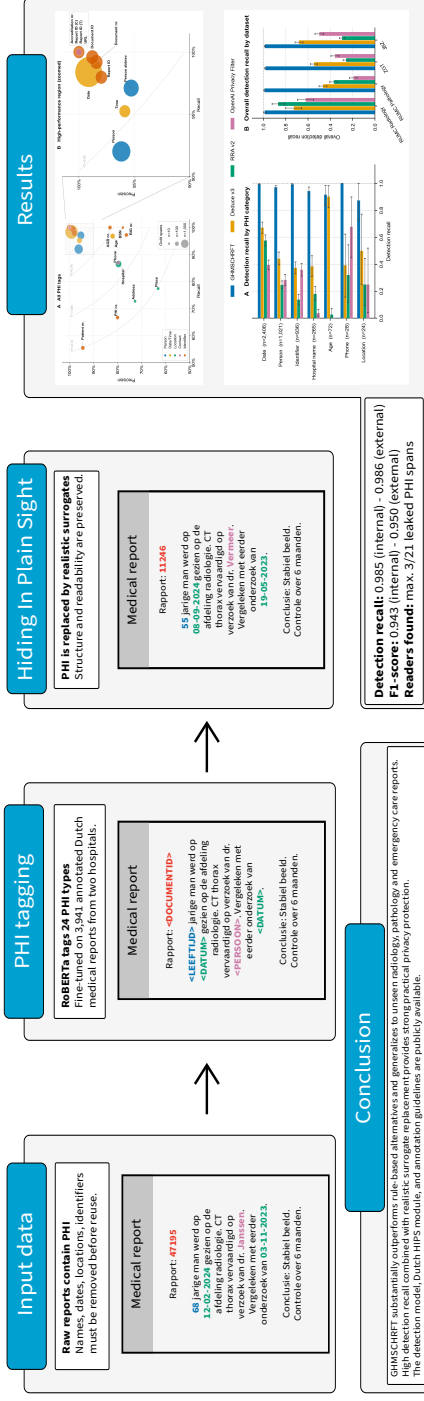


Graphical Abstract

GHMSCHRFT: An AI-Based De-identification Pipeline for Dutch Clinical Reports

Luc Builtjes, Joeran S. Bosma, Judith Lefkes, Anouk Veldhuis, Jeroen Geerdink, Tijmen de Haas, Fabian R. Hoitsma, Rianne A. Weber, Sarah de Boer, Myrthe Buser, Harry J. Kendall, Delano Sanches, Bram van Ginneken, Henkjan Huisman, Ewoud J. Smit, Mathias Prokop, Alessa Hering

The GHMSCHRFT Dutch medical text de-identification pipeline



GHMSCHRFT: An AI-Based De-identification Pipeline for Dutch Clinical Reports

Luc Builtjes^{a,d,*}, Joeran S. Bosma^a, Judith Lefkes^b, Anouk Veldhuis^c, Jeroen Geerdink^c, Tijmen de Haas^a, Fabian R. Hoitsma^a, Rianne A. Weber^a, Sarah de Boer^a, Myrthe Buser^a, Harry J. Kendall^a, Delano Sanches^a, Bram van Ginneken^a, Henkjan Huisman^a, Ewoud J. Smit^a, Mathias Prokop^a, Alessa Hering^a

^a*Department of Medical Imaging, Radboudumc, Geert Grooteplein Zuid 10, Nijmegen, 6525GA, The Netherlands*

^b*Department of Pathology, Radboudumc, Geert Grooteplein Zuid 10, Nijmegen, 6525GA, The Netherlands*

^c*Department of Information Management, Ziekenhuisgroep Twente, Zilvermeeuw 1, Almelo, 7609PP, The Netherlands*

^d*Department of Medical Imaging, Jeroen Bosch Ziekenhuis, Deutersestraat 2, 's-Hertogenbosch, 5223GZ, The Netherlands*

Abstract

Objective: Secondary use of clinical free-text requires reliable de-identification, but existing open-source methods do not yet perform well enough for practical deployment without institution specific tuning. We aimed to develop and validate GHMSCHRFT (from Dutch *geheimschrift*, ‘cipher’), an open-source de-identification pipeline for Dutch clinical text, and to assess whether its protection holds up to human scrutiny after deployment.

Methods: GHMSCHRFT combines a transformer-based model (RoBERTa-large, pretrained on Dutch medical text) fine-tuned for PHI detection on 3,941 annotated reports from two hospitals, spanning radiology, pathology, and multiple outpatient specialties, with a Dutch-specific Hiding in Plain Sight (HIPS) surrogate replacement module. Performance was evaluated on internal (980 reports) and external (340 reports from an unseen institution) test sets, and a reader study assessed residual disclosure risk by having five readers attempt to identify leaked PHI in HIPS-processed output.

*Corresponding author

Email address: luc.builtjes@radboudumc.nl (Luc Builtjes)

Results: GHMSCHRFT achieved detection recall of 0.985 and 0.986, with precision of 0.941 and 0.949, on internal and external test sets, compared with recall of 0.549/0.682 for DEDUCE, 0.386/0.295 for an in-house rule-based system, and 0.339/0.498 for the OpenAI Privacy Filter. In the reader study, individual detection rates ranged from 0–14.3% of evaluable leaked PHI spans (5 of 21 identified collectively), with re-identification precision no higher than 8.3%.

Conclusion: GHMSCHRFT generalizes across institutions for radiology, pathology, and emergency care reports, outperforming rule-based alternatives by more than 30 percentage points in detection recall. On top of these detection results, the reader study confirms that the full pipeline provides strong practical protection: readers performed no better than chance at distinguishing genuine leaks from correctly replaced surrogates. The detection model, Dutch HIPS module, and annotation guidelines are publicly available.

Keywords: de-identification, protected health information, clinical natural language processing, hiding-in-plain-sight, GDPR

1. Introduction

Electronic health records typically contain two broad categories of data: structured fields such as diagnosis codes and date of birth, and unstructured free-text reports created by clinicians in natural language. Structured data are directly machine-interpretable, but much of the clinically relevant information is expressed in narrative text [1, 2]. Free-text reports capture longitudinal reasoning, temporal relations between findings, and contextual nuance that structured fields cannot easily represent [3]. This makes clinical narratives a valuable resource for secondary use, especially for the development of data-driven clinical tools [4].

Realizing this potential requires de-identifying clinical narratives before patient data can be shared and reused responsibly. Clinical narratives routinely contain protected health information (PHI): direct and quasi-identifiers such as person names, dates, locations, and institutional identifiers that can be used to re-identify individuals. In the European Union, the General Data Protection Regulation (GDPR) governs the processing of such personal data [5]. De-identification, the process of detecting and removing or replacing PHI to a standard that is robust against reasonably likely means of re-identification,

19 is therefore a prerequisite for compliant data sharing and secondary use of
20 clinical text.

21 Automated de-identification has been an active research area for several
22 decades. Early systems relied on handcrafted rules, dictionaries, and regular
23 expressions to detect identifiers [6]. These approaches work well for struc-
24 turally regular identifiers such as phone numbers or postal codes, but require
25 substantial manual effort to develop, are fragile under format variation, and
26 tend to generalize poorly when applied outside the institution for which they
27 were designed. The introduction of large annotated corpora and pretrained
28 language models shifted the field toward data-driven approaches. Transformer-
29 based sequence-labeling models now represent the dominant paradigm and
30 consistently outperform rule-based systems when sufficient training data are
31 available [7, 8]. Generative large language models (LLMs) have more recently
32 been explored as an alternative to supervised sequence labeling [9, 10]. These
33 prompt-based approaches avoid the need for annotated training data, but
34 in practice require careful prompt engineering and are sensitive to phrasing,
35 output format, and model version, making reliable deployment difficult [11].
36 Fine-tuning generative LLMs on annotated PHI data can improve consistency,
37 but introduces a separate privacy concern: language models are known to
38 memorize fragments of their training data, and a publicly released fine-tuned
39 model may regurgitate patient information seen during training [12]. This
40 tension between performance and safe release makes generative LLMs a poor
41 fit for de-identification pipelines intended for open distribution.

42 On top of this, both supervised and generative approaches have been
43 developed and evaluated predominantly on English clinical text [13], and
44 it remains unclear how either transfers to other languages or institutional
45 settings. A recent example is the OpenAI Privacy Filter [14], an open-source
46 model developed specifically for de-identification but trained on general,
47 primarily English text, which illustrates both the progress and the remaining
48 language gap.

49 For Dutch clinical text, the landscape looks markedly different; Dutch
50 hospitals have largely developed de-identification solutions independently,
51 but little of this work has been released openly. The main exception is
52 DEDUCE [15], a rule-based system developed for de-identification of Dutch
53 clinical notes, which has become the de facto open-source standard in the
54 Netherlands. No publicly available transformer-based alternative exists for
55 Dutch. This situation is compounded by the difficulty of sharing clinical data
56 across institutions under GDPR, which limits the assembly of the multicenter

57 annotated corpora needed to train and rigorously evaluate learning-based
58 approaches. As a result, it remains unclear whether transformer models
59 trained on Dutch clinical data from one or more institutions can achieve the
60 detection performance and cross-institutional generalizability required for
61 practical deployment.

62 Detection performance alone, however, does not fully determine privacy
63 risk. A common approach after detection is to mask or delete identified
64 spans. This leaves the document structure intact but creates an unambiguous
65 signal: any identifier that escapes detection remains in plain text while
66 everything around it is redacted, making residual PHI easy to spot. Hiding
67 in Plain Sight (HIPS) was proposed as an alternative strategy to address
68 this vulnerability [16]. Rather than removing detected identifiers, HIPS
69 replaces them with realistic surrogates, so that any residual PHI blends into a
70 document that still reads as natural clinical text. Carrell et al. demonstrated
71 that this approach meaningfully reduces the detectability of leaked identifiers
72 by both automated systems and human reviewers [17]. Despite this evidence,
73 HIPS implementations remain rare and almost exclusively English. No open-
74 source Dutch medical module exists. Naively applying English surrogate
75 generators to Dutch text would produce implausible output, since Dutch
76 name distributions, address formats, and institutional identifier structures
77 differ substantially from their English counterparts. This undermines the
78 core idea of HIPS: residual PHI would stand out against foreign-language
79 surrogates rather than blend in, becoming identifiable rather than protected.

80 Together, these gaps raise two related questions. Can a transformer model
81 trained on multicenter Dutch clinical data achieve detection performance
82 sufficient for practical deployment, including generalization to hospitals not
83 seen during training? And does replacing detected identifiers with Dutch-
84 specific surrogates reduce residual disclosure risk to a degree that is empirically
85 verifiable through human assessment?

86 In this study, we address both questions by presenting and evaluating GHM-
87 SCHRFT (from Dutch *geheimschrift*, ‘cipher’), an end-to-end de-identification
88 pipeline for Dutch clinical free-text reports. We fine-tune a Dutch clinical
89 transformer model on multicenter annotated data from two hospitals and
90 evaluate detection performance on internal and external test sets. To address
91 the absence of Dutch-specific surrogate replacement tooling, we developed
92 **dutch-med-hips** [18], an open-source HIPS module that replaces detected
93 PHI with realistic Dutch surrogates. We assess whether the pipeline pro-
94 vides sufficient protection in practice, by running it end-to-end on a mixed

95 set of reports, where some contain residual identifiers and others are fully
96 de-identified. We then ask human readers to assess whether any remaining
97 PHI can be detected. Together, these components constitute GHMSCHRFT:
98 a language-specific, practically oriented, and empirically verified framework
99 for the trustworthy secondary use of Dutch clinical text.

100 The main contributions of this work are:

- 101 • A publicly available, fine-tuned Dutch clinical transformer model for
102 PHI detection, evaluated on both internal and external test sets to
103 assess cross-institution generalizability.
- 104 • A detailed annotation protocol with 24 granular categories for Dutch
105 medical PHI, including span-merging and edge-case rules, provided in
106 the Supplementary Material to support reproducible annotation.
- 107 • An open-source Dutch hiding-in-plain-sight surrogate replacement mod-
108 ule (<https://github.com/DIAGNijmegen/dutch-med-hips>) providing
109 type-consistent surrogate generation.
- 110 • A human reader study evaluating residual re-identification risk, provid-
111 ing empirical evidence of the practical protection offered by the complete
112 end-to-end system.

Statement of Significance

Problem or Issue	Secondary use of Dutch clinical free-text is severely limited by the lack of reliable de-identification tools. Existing open-source systems were developed for English and achieve insufficient detection recall for practical deployment.
What is Already Known	Transformer-based models outperform rule-based de-identification systems when annotated training data are available. Surrogate replacement (HIPS) better preserves data utility than simple masking. No validated open-source transformer pipeline exists for Dutch clinical text.
What this Paper Adds	GHMSCHRFT is the first open-source transformer-based de-identification pipeline for Dutch clinical reports. It achieves detection recall of 0.985–0.986 across internal and external test sets, exceeding existing tools by more than 30 percentage points. A reader study confirms that the full pipeline provides strong practical protection. The model, surrogate module, and annotation guidelines are publicly released.
Who Would Benefit	Clinical researchers and hospital data governance teams seeking to enable secondary use of Dutch clinical free-text, and the broader clinical NLP community developing infrastructure for non-English health data.

113 2. Methods

114 Figure 1 provides a high-level overview of the GHMSCHRFT pipeline.
 115 The following subsections describe each component in detail.

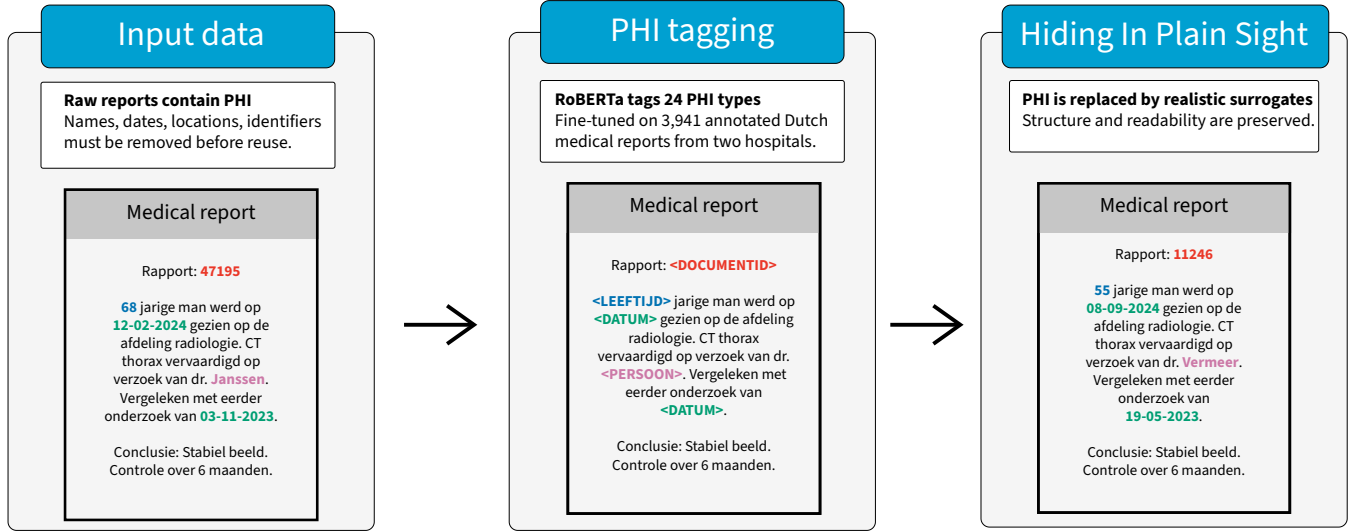


Figure 1: Overview of the GHMSCHRFT de-identification pipeline. Raw Dutch clinical reports containing PHI (left) are processed by a fine-tuned RoBERTa model that detects and tags 24 PHI categories (center). Tagged PHI is then replaced with realistic surrogates (right), preserving the report’s structure and readability while rendering any missed PHI indistinguishable from the surrogates.

116 2.1. Data Sources and Annotation

117 We conducted a retrospective multicenter study to develop and evalu-
 118 ate GHMSCHRFT, an automated de-identification framework for Dutch
 119 clinical free-text reports. Data came from three hospitals: Radboud Univer-
 120 sity Medical Center (RUMC), Ziekenhuisgroep Twente (ZGT), and Jeroen
 121 Bosch Ziekenhuis (JBZ). Development used RUMC and ZGT only: RUMC
 122 contributed manually annotated radiology and pathology reports; ZGT con-
 123 tributed annotated clinical and outpatient letters across eleven specialties
 124 (internal medicine, cardiology, rheumatology, pulmonology, pediatrics, surgery,
 125 allergology, gynecology, emergency medicine, gastroenterology, and urology).

126 JBZ contributed radiology, pathology, and emergency care reports, and was
 127 used for external validation only.

128 All processing was performed in secure institutional environments under
 129 GDPR and local governance policies, with ethical approval from relevant
 130 review boards. To enable GDPR-compliant cross-institutional transfer, a
 131 preliminary RUMC-trained model automatically tagged PHI at ZGT and
 132 JBZ. Annotators then corrected predictions in Doccano [19]. Prior to transfer,
 133 detected PHI spans were replaced with HIPS surrogates, preserving clinical
 134 structure while minimizing disclosure risk. Resulting annotations were stored
 135 as character-offset spans mapped to the surrogate-replaced text.

136 Annotation quality was ensured through a two-stage review process. Mul-
 137 tiple readers independently annotated each report in a first pass. A subset
 138 of reports was doubly annotated to assess inter-reader agreement, quantified
 139 using character-level Cohen’s kappa restricted to PHI positions. An author
 140 subsequently reviewed all annotations, resolving ambiguities and correcting er-
 141 rors. Annotation rules followed the PHI guidelines in Supplementary Material
 142 (Section S7).

143 Dataset statistics are summarised in Table 1. The internal test sets
 144 comprised 741 RUMC and 239 ZGT reports, while all 340 JBZ reports served
 145 as an external test set, excluded from all development. The remaining 3,941
 146 reports formed the development set. Fine-grained labels were represented in
 147 Beginning-Inside-Outside (BIO) format across 24 training categories. The
 148 complete tag inventory is provided in Supplementary Table S1.

Table 1: Dataset statistics by hospital and split. RUMC: Radboud University Medical Center (Radboudumc, Nijmegen). ZGT: Ziekenhuisgroep Twente (Almelo/Hengelo). JBZ: Jeroen Bosch Ziekenhuis (’s-Hertogenbosch). Development reports were used for training and validation. HIPS-augmented duplicates and 50 synthetic reports were added during training but are not counted here. Average entities per report is shown in parentheses.

	Development		Internal test		External test	
Dataset	Reports	Entities	Reports	Entities	Reports	Entities
RUMC	2,983	8,731 (2.9)	741	2,224 (3.0)	—	—
ZGT	958	10,684 (11.2)	239	2,675 (11.2)	—	—
JBZ	—	—	—	—	340	1,818 (5.3)
Total	3,941	19,415 (4.9)	980	4,899 (5.0)	340	1,818 (5.3)

149 *2.2. PHI Detection Model Development and Training Augmentation*

150 PHI detection was trained by fine-tuning `dragon-roberta-large-mixed-domain` [20], a RoBERTa-large model pretrained on mixed-domain
151 Dutch medical text, using the DRAGON baseline framework [21] with a
152 token-classification NER head. Input sequences were left-truncated to 512
153 tokens during training. At inference, documents exceeding 512 tokens are
154 processed using a strided sliding-window approach to ensure full-document
155 coverage. Training used HuggingFace Transformers with a learning rate of
156 1×10^{-5} , effective batch size 8 (batch size 1, 8 gradient-accumulation steps),
157 and gradient checkpointing; the best checkpoint was selected by validation
158 performance.
159

160 Two augmentation strategies were applied to the training split. First, 50
161 fully synthetic reports were generated with `qwen3:30b-a3b-thinking-2507-q4_K_M` [22]
162 to address class imbalance for rare identifier types (e.g., IBAN, BSN).
163 Second, HIPS augmentation via `dutch-med-hips` [18] produced an
164 additional surrogate-replaced version of each annotated report. The final
165 training set combined original development reports, their HIPS-augmented
166 counterparts, and the 50 synthetic reports.

167 *2.3. Dutch HIPS Surrogate Module (`dutch-med-hips`)*

168 `dutch-med-hips` [18] replaces tagged PHI with type-consistent surrogates,
169 covering names and initials, ages, dates and times, hospital and location names,
170 study names, contact details, and structured identifiers (BSN, IBAN, report
171 IDs, institution-specific numbers). Surrogate generation uses a Dutch-locale
172 fake data library, sampling predominantly Dutch names with a lower propor-
173 tion of foreign names to reflect realistic patient demographics, supplemented by
174 curated geographic lookup pools and template-driven synthesis for structured
175 identifiers. The engine supports deterministic behavior via document-hash
176 seeding and returns de-identified text with a structured original-to-surrogate
177 mapping including character offsets.

178 At inference, GHMSCHRFT applies four sequential steps: text prepro-
179 cessing, PHI detection, post-processing, and surrogate replacement. Post-
180 processing merges adjacent same-type predictions to reduce span fragmenta-
181 tion from subword tokenization, after which finalized spans are replaced by
182 the HIPS module.

183 2.4. Evaluation Metrics

184 Performance was evaluated using two primary metrics. Micro-averaged
185 entity-level F1-score (exact span and type match, via `sequeval`/DragonEval
186 [21]) was the principal summary measure. Detection recall served as a safety-
187 oriented complement: a PHI entity was counted as detected when every
188 token in the gold span received any non-0 label, regardless of type or span
189 extent, making false negatives under this metric a direct proxy for residual
190 disclosure risk. Macro-averaged F1-score and per-category breakdowns are
191 additionally reported. All headline metrics use 95% confidence intervals from
192 document-level bootstrap resampling (5,000 iterations). As a measure of
193 inadvertent information loss, we counted predicted PHI spans that had no
194 overlap with any annotated gold span; these represent clinical content that
195 the pipeline would replace with a surrogate despite not being PHI.

196 2.5. Baseline Comparison

197 GHMSCHRFT was compared against DEDUCE v3.0.6 [15] (the main
198 open-source rule-based system for Dutch clinical text), the OpenAI Privacy
199 Filter [14] (a general-purpose neural de-identification model), and the Rad-
200 boudumc Radiology Report Anonymizer (RRA) v2.14 (an institution-specific
201 in-house comparator). Because these systems use only partially overlapping
202 tag inventories, comparisons were conducted across seven compiled cate-
203 gories with clear cross-system counterparts (Supplementary Table S7), using
204 compiled-category detection recall on both test sets.

205 2.6. End-to-End Residual Risk Check

206 To assess residual disclosure risk in the final deployment scenario, five
207 native Dutch-speaking readers (three clinical researchers, two radiology resi-
208 dents) independently reviewed 100 HIPS-processed de-identified reports in
209 Doccano [19]. Readers marked fragments believed to be original PHI missed
210 by the pipeline as either definite or suspected residual PHI. Readers were
211 instructed not to flag correctly replaced surrogates and were blinded to report
212 composition.

213 The 100 reports comprised the complete set of cases with at least one
214 pipeline-missed PHI span in the combined internal test set ($n = 50$, exhausting
215 all error cases), and an equal number of fully de-identified cases drawn from
216 the same test set. Prior to the study, one author reviewed all 50 error cases
217 to assess human detectability in context. Thirty-two were reclassified as
218 effectively de-identified because the remaining fragments were undetectable

219 after partial surrogate replacement (e.g., a date reduced to a bare digit).
 220 This left 18 cases with 21 evaluable leaked spans; combined with the 50
 221 fully de-identified and 32 reclassified cases, 82 of 100 reports were treated
 222 as correctly de-identified. Reader performance was quantified by counting
 223 correctly identified leaked fragments (reader annotation overlapping a gold
 224 leaked span), missed fragments, and false positives. Exact 95% confidence
 225 intervals for per-reader detection rates were computed using the Clopper–
 226 Pearson method. To quantify surrogate dilution, the precision upper bound
 227 for a perfect PHI-shape detector was computed as $L/(L + S)$, where L is
 228 the total number of evaluable leaked spans and S the total number of HIPS
 229 surrogate replacements across the 100 documents. A per-document Monte
 230 Carlo simulation (10,000 iterations, seed 42) compared each reader’s observed
 231 true-positive count against the null distribution expected under uniform
 232 random sampling of the same annotation volume from the $L + S$ candidate
 233 pool.

234 3. Results

235 3.1. Inter-reader Agreement

236 Inter-reader agreement was high in both doubly annotated subsets. In
 237 RUMC radiology, character-level Cohen’s kappa restricted to PHI positions
 238 was 0.940. In RUMC pathology, character-level Cohen’s kappa restricted to
 239 PHI positions was 0.967.

240 3.2. Overall Internal Performance

241 On the combined internal test set of 980 reports containing 4,899 annotated
 242 PHI entities, the model achieved micro-averaged precision, recall, and F1-score
 243 of 0.941, 0.944, and 0.943, respectively. Detection evaluated independently of
 244 label correctness yielded precision, recall, and F1-score of 0.948, 0.950, and
 245 0.949. Performance was strongest for structured identifier categories, with
 246 accreditation numbers, BIG registration numbers, document identifiers, and
 247 URLs all achieving F1-scores of 0.99 or higher. Lower F1-scores were observed
 248 for rarer or semantically broader categories, including company names (0.25),
 249 study names (0.55), addresses (0.55), and place names (0.62). Among the
 250 most frequent categories, F1-scores were 0.97 for dates, 0.89 for person names,
 251 0.95 for initials, 0.90 for times, and 0.81 for hospital names.

252 3.3. External Validation

253 On the JBZ external test set of 340 reports containing 1,818 annotated PHI
 254 entities, merged-entity evaluation achieved micro-averaged precision, recall,
 255 and F1-score of 0.949, 0.952, and 0.950, respectively. Detection evaluated
 256 independently of label correctness yielded precision, recall, and F1-score of
 257 0.981, 0.984, and 0.983. Performance patterns were broadly consistent with
 258 the internal results; per-category breakdowns are provided in Table 2.

259 3.4. Per-tag Performance

Table 2: Per-tag performance aggregated across the internal test sets (RUMC radiology, RUMC pathology, ZGT; $N = 980$ reports) and the external JBZ test set ($N = 340$ reports), evaluated under merged-entity scoring. Detection recall counts a gold entity as detected when every token in the annotated span receives a non-0 label, irrespective of assigned entity type or predicted span extent. Precision, recall, and F1 require the exact correct entity label and span boundaries. The **Overall** row reports micro-averaged metrics; 95% bootstrap confidence intervals (5,000 iterations, document-level resampling) are shown in parentheses on the following line. Per-dataset breakdowns are provided in Supplementary Table S2.

Tag	N	Det. rec.	Prec.	Rec.	F1
Internal test set (RUMC & ZGT; $N = 980$)					
accreditatie nummer (<i>accreditation no.</i>)	111	1.00	1.00	1.00	1.00
adres (<i>address</i>)	11	0.91	0.55	0.55	0.55
agbnummer (<i>AGB no.</i>)	42	1.00	0.98	1.00	0.99
bedrijf (<i>company</i>)	2	0.50	0.17	0.50	0.25
bignummer (<i>BIG no.</i>)	35	1.00	1.00	1.00	1.00
bsn (<i>citizen no.</i>)	7	1.00	0.78	1.00	0.88
datum (<i>date</i>)	2406	1.00	0.98	0.97	0.97
documentid (<i>document ID</i>)	245	1.00	0.99	1.00	0.99
documentnummer (<i>document no.</i>)	252	1.00	0.99	0.99	0.99
leeftijd (<i>age</i>)	72	0.92	0.73	0.92	0.81
patientnummer (<i>patient no.</i>)	11	0.91	0.90	0.82	0.86
persoon (<i>person</i>)	592	0.96	0.91	0.87	0.89
persoonafkorting (<i>initials</i>)	429	0.98	0.94	0.95	0.95
phinummer (<i>other PHI no.</i>)	35	0.97	0.75	0.60	0.67

Continued on next page

Table 2 (continued)

Tag	N	Det. rec.	Prec.	Rec.	F1
plaats (<i>place</i>)	13	0.85	0.56	0.69	0.62
rapport id (<i>report ID</i>)	198	0.98	0.97	0.97	0.97
studie naam (<i>study name</i>)	8	1.00	0.43	0.75	0.55
telefoonnummer (<i>phone no.</i>)	28	1.00	0.74	0.89	0.81
tijd (<i>time</i>)	129	0.92	0.91	0.90	0.90
url (<i>URL</i>)	8	1.00	1.00	1.00	1.00
ziekenhuis (<i>hospital</i>)	265	0.94	0.78	0.83	0.81
Overall	4899	0.983	0.941	0.944	0.943
		(0.98, 0.99)	(0.93, 0.95)	(0.93, 0.95)	(0.93, 0.95)
External test set (JBZ; $N = 340$)					
accreditatie nummer (<i>accreditation no.</i>)	100	1.00	1.00	1.00	1.00
bignummer (<i>BIG no.</i>)	3	1.00	0.04	0.33	0.08
datum (<i>date</i>)	807	0.99	0.99	0.99	0.99
leeftijd (<i>age</i>)	78	0.99	0.94	0.99	0.96
patientnummer (<i>patient no.</i>)	70	1.00	0.97	0.47	0.63
persoon (<i>person</i>)	482	0.97	0.98	0.95	0.96
persoonafkorting (<i>initials</i>)	77	1.00	0.94	0.99	0.96
phinummer (<i>other PHI no.</i>)	2	1.00	0.00	0.00	0.00
plaats (<i>place</i>)	5	0.80	0.67	0.80	0.73
rapport id.c nummer (<i>report ID C-no.</i>)	41	1.00	1.00	1.00	1.00
rapport id.t nummer (<i>report ID T-no.</i>)	62	1.00	1.00	1.00	1.00
telefoonnummer (<i>phone no.</i>)	21	1.00	0.94	0.71	0.81
tijd (<i>time</i>)	58	0.97	0.97	0.97	0.97
ziekenhuis (<i>hospital</i>)	12	1.00	0.48	0.92	0.63
Overall	1818	0.986	0.949	0.952	0.950
		(0.98, 0.99)	(0.94, 0.96)	(0.94, 0.96)	(0.94, 0.96)

Figure 2 summarizes span-level precision and recall by PHI tag across all datasets and highlights that most frequent tag types cluster in the high-performance region, whereas lower-performing tags are generally rare and semantically broader.

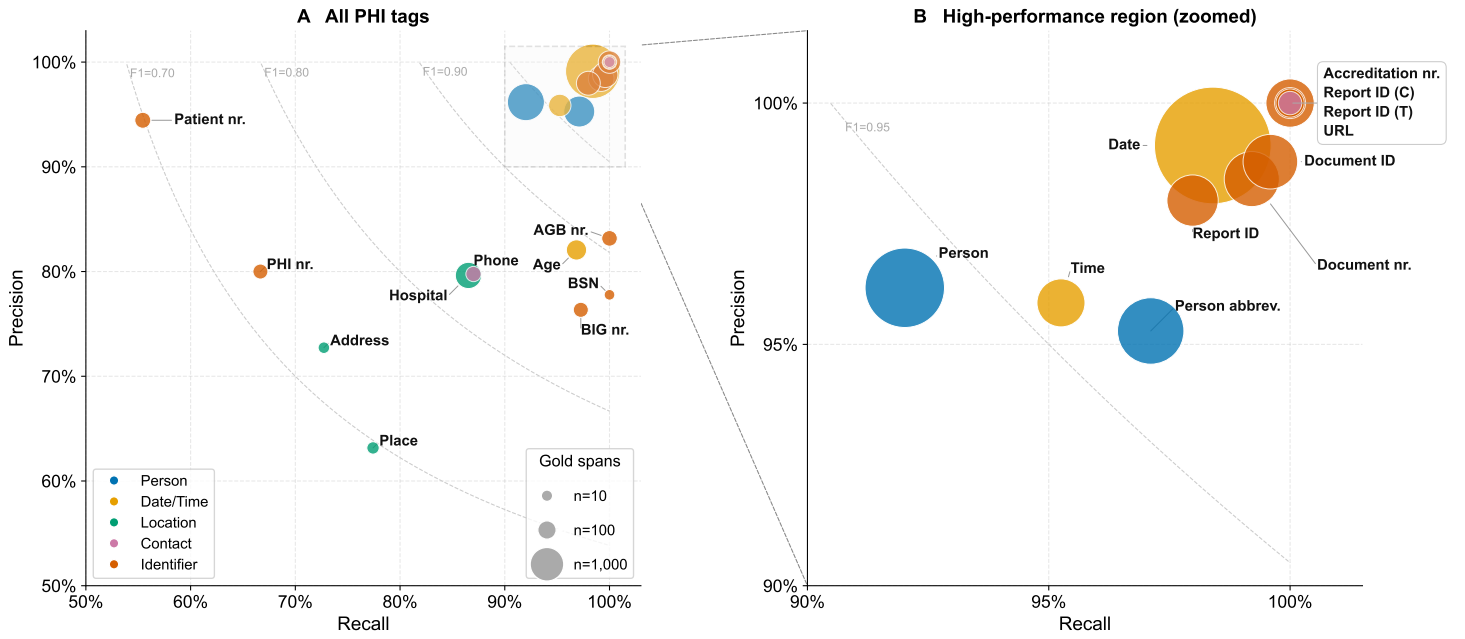


Figure 2: Span-level precision and recall per PHI tag, aggregated across all datasets. Each bubble represents one tag type; bubble area is proportional to the number of gold spans (n). Tags are coloured by PHI category. F1 iso-contours are shown as dashed lines. Both axes start at 0.5 rather than 0. (A) All tag types across the full precision–recall space. The dashed box indicates the region expanded in panel B. (B) Zoomed view of the high-performance region ($\geq 90\%$ precision and recall). Tags that overlap at identical coordinates (100%, 100%) are listed in the inset. Span matching is exact: a predicted span is counted as correct only if its token boundaries and tag type exactly match the gold annotation.

264 3.5. Comparison with Rule-based Baselines

265 GHMSCHRFT achieved an overall compiled-category detection recall
 266 of 0.985 (95% CI 0.98–0.99) on the combined internal test data, compared
 267 with 0.549 (0.53–0.57) for DEDUCE, 0.386 (0.36–0.41) for RRA, and 0.339
 268 (0.31–0.36) for the OpenAI Privacy Filter. On the external JBZ test set,
 269 GHMSCHRFT retained strong performance with an overall recall of 0.986
 270 (0.98–0.99) and substantially outperformed DEDUCE (0.682, 0.65–0.72),
 271 RRA (0.295, 0.27–0.32), and the OpenAI Privacy Filter (0.498, 0.46–0.54).
 272 Per-category and per-dataset breakdowns are provided in Supplementary
 273 Table S5.

274 Across the internal datasets, GHMSCHRFT showed the largest gains
 275 for person names, dates, hospital names, and compiled identifier categories.
 276 DEDUCE remained competitive for age expressions and some date-heavy
 277 subsets, whereas RRA performed best on the in-domain radiology subsets
 278 but generalized poorly to ZGT and JBZ. The OpenAI Privacy Filter showed
 279 moderate recall for person names and dates in radiology subsets but near-
 280 zero recall for hospital names and age expressions across all datasets, likely
 281 reflecting a domain and language mismatch with Dutch clinical text. Figure 3
 282 visualizes these category-level and dataset-level differences.

283 3.6. False-Positive Rate and Information Loss

284 Of 6,734 predicted PHI spans across both test sets, 6,353 (94.3%) were
 285 exact true positives, 271 (4.0%) overlapped a gold span but with incorrect
 286 boundaries or tag type, and 110 (1.6%) had no overlap with any gold span.
 287 These 110 pure false positives comprised 132 tokens across 1,320 reports
 288 (0.08 per report), representing clinical content that would be replaced with a
 289 surrogate in a deployed system. Full per-tag results are provided in Supple-
 290 mentary S3.

291 3.7. End-to-End Residual Risk Assessment

292 After applying the exclusions described in the Methods, 18 of the 100
 293 presented cases contained evaluable residual PHI, comprising 21 leaked PHI
 294 spans in total. Per-reader detection rates with exact 95% confidence intervals
 295 are reported in Table 3. No reader identified more than three spans: DS
 296 identified 3 (14.3%; 95% CI 3.0–36.3%), RW and HK each identified 2 (9.5%;
 297 1.2–30.4%), SB identified 1 (4.8%; 0.1–23.8%), and MB identified none (0.0%;
 298 0.0–16.1%). Confidence intervals are wide due to the small number of evaluable
 299 spans. Collectively, five of the 21 leaked spans were identified by at least

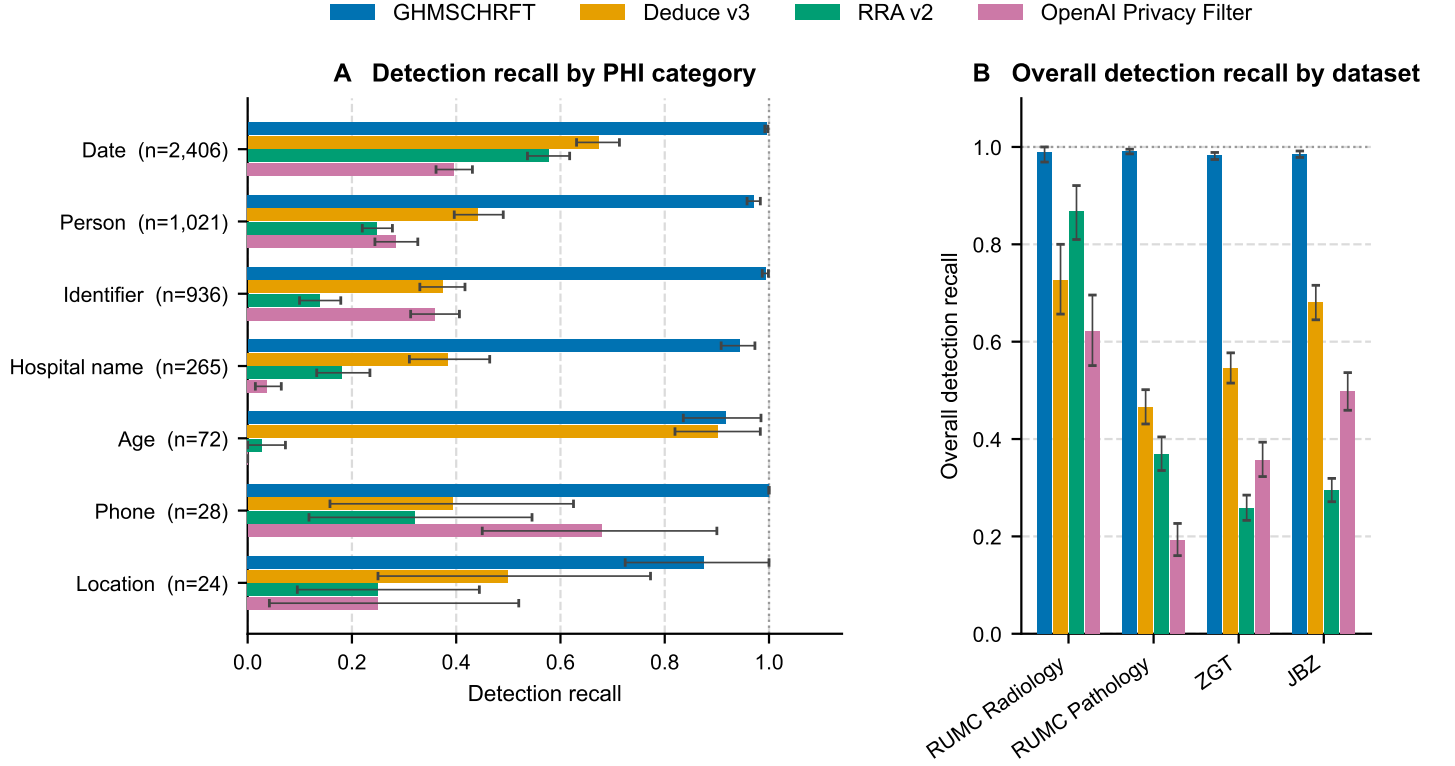


Figure 3: Detection recall per PHI category and per dataset for GHMSCHRFT and three comparator systems. A span is counted as detected if every token in the gold span receives a non-0 prediction, irrespective of assigned label. (A) Per-category detection recall on the combined internal test set; n denotes the number of gold spans per category. Categories are ordered by frequency. (B) Overall detection recall per system across four hospital datasets. GHMSCHRFT was trained on RUMC and ZGT data; JBZ is an unseen external institution. DEDUCE v3 and RRA v2 are rule-based Dutch clinical de-identification systems. The OpenAI Privacy Filter is a 1.5B sparse mixture-of-experts model trained on general, primarily English text. Detection recall is reported rather than F1 to reflect the safety-critical requirement that no PHI is missed. Error bars denote 95% bootstrap confidence intervals (5,000 iterations, document-level resampling).

one reader: a partially-leaked surname, an age fragment, a hospital panel abbreviation, and a first name appearing twice in the same document (each occurrence counted as a separate span). The remaining 16 spans went undetected by every reader.

Readers generated substantially more annotations that did not correspond to any leaked PHI (false positives). RW made 57 annotations in total, of which 55 were false positives; DS made 36, of which 34 were false positives; MB made 34, all false positives; SB made 33, of which 32 were false positives; and HK made 25, of which 24 were false positives. Across readers, 0–28% of flagged error cases included at least one ‘definite PHI’ annotation. The remainder were labeled ‘suspected PHI’. False-positive annotations arose under both confidence levels: several readers definitively flagged HIPS-generated surrogate names as leaked PHI while overlooking actual residual content.

A surrogate dilution analysis quantified the structural precision limit. The 100 reader-study documents contained $S = 1,039$ HIPS surrogate replacements alongside $L = 21$ evaluable leaked spans, yielding a combined PHI-shaped pool of 1,060 candidates. The theoretical precision ceiling for a perfect PHI-shape detector on this dataset is $L/(L + S) = 2.0\%$. All observed reader precisions (0%–8.3%) fell within the 95% interval of the per-document Monte Carlo simulation, confirming that no reader performed significantly above the level expected from random selection within this pool.

Table 3: Span-level results per reader in the end-to-end residual risk assessment. The 18 evaluable error cases contained 21 leaked PHI spans in total. Detection rate is the fraction of leaked PHI spans overlapped by at least one reader annotation; 95% confidence intervals are exact Clopper–Pearson intervals. False positives (FP) are reader annotations not overlapping any gold leaked fragment. Spans found may exceed TP reader annotations when a single reader annotation simultaneously overlapped two adjacent leaked spans.

Reader	Role	Spans found	Detection rate (95% CI)	FP
RW	Clinical researcher	2/21	9.5% (1.2–30.4%)	55
SB	Clinical researcher	1/21	4.8% (0.1–23.8%)	32
MB	Clinical researcher	0/21	0.0% (0.0–16.1%)	34
HK	Radiology resident	2/21	9.5% (1.2–30.4%)	24
DS	Radiology resident	3/21	14.3% (3.0–36.3%)	34

321 4. Discussion

322 GHMSCHRFT provides an end-to-end Dutch de-identification pipeline
323 with high detection recall, external generalization, and practical residual-risk
324 evaluation, addressing a gap that has persisted since rule-based systems
325 became the de facto standard. Near-identical performance on the internal and
326 external JBZ test sets on radiology, pathology, and emergency care reports
327 suggests the model has learned generalizable representations of Dutch clinical
328 PHI rather than institution-specific formatting conventions for these report
329 types. ZGT letters were part of internal development, so the evidence for
330 generalization in that setting is within-institution only. This likely reflects
331 the relevance of the pretraining domain: Bosma et al. [21] demonstrated that
332 Dutch medical pretraining consistently improves clinical NLP performance,
333 and the present results extend that finding to de-identification.

334 GHMSCHRFT’s recall advantage over the rule-based baselines is best
335 illustrated by RRA. Built specifically for RUMC radiology reports, RRA
336 achieved high recall on that subset but dropped markedly on ZGT and JBZ,
337 and notably, even on RUMC pathology reports, despite originating from the
338 same institution. This generalization failure is typical of rule-based systems,
339 which struggle whenever reporting templates or documentation conventions
340 differ from those they were built for. DEDUCE showed a similar, though less
341 pronounced, shortfall. It was, however, evaluated without institution-specific
342 tuning. Manual fine-tuning per target institute could improve its performance.

343 OpenAI Privacy Filter’s poor performance has a different explanation.
344 As a neural model trained on general, primarily English text, its near-zero
345 recall for hospital names and age expressions reflects a mismatch with Dutch
346 clinical text rather than an architectural limitation. To probe this further, we
347 additionally evaluated a version fine-tuned on English medical text [23]. Full
348 results are provided in Supplementary S5. This fine-tuned variant did not
349 outperform the out-of-the-box model on our Dutch data. As its fine-tuning
350 remained in English, it addressed neither the Dutch language itself nor the
351 Netherlands-specific institutional and identifier conventions that distinguish
352 Dutch clinical documentation.

353 At a more granular level, per-tag performance was strongest for structurally
354 well-defined identifier types. Lower F1-scores for company names (0.25),
355 study names (0.55), addresses (0.55), and place names (0.62) are primarily
356 attributable to low training support rather than inherent difficulty, as the
357 uniformly high detection recall shows the model generally identified these

358 spans as PHI. Errors were more often misclassification or boundary mistakes
359 than outright missed detections. For addresses, errors were almost exclusively
360 boundary errors. For example, a full street address would be correctly detected
361 with only a trailing province abbreviation excluded from the predicted span.

362 Two related anomalies appeared in the external test set: patient numbers
363 achieved perfect detection recall but an F1 of only 0.63, while BIG registration
364 numbers showed precision of just 0.04 despite only three gold instances.
365 Manual inspection revealed a shared cause. The model consistently detected
366 patient number spans but misclassified them as BIG numbers, because JBZ
367 uses a numeric patient identifier format superficially resembling BIG numbers
368 not seen during training. This is a cross-institution format mismatch, not a
369 detection failure.

370 A separate, more critical failure mode appeared primarily in clinical letters,
371 where patients are occasionally referred to by first name only. This pattern
372 was absent from radiology and pathology reports that dominated training.
373 The model had limited exposure to this and did not reliably recognize isolated
374 first names as PHI in clinical correspondence. Future work will address this
375 by including additional annotated examples.

376 The false-positive analysis (Supplementary S3) indicates that the pipeline
377 accidentally destroys very little useful clinical information. Only 110 of 6,734
378 predicted spans (1.6%) had no overlap with any gold PHI annotation, affecting
379 132 tokens across 1,320 reports (0.08 per case). Hospital names (ZIEKEN-
380 HUIS, 27 cases) and age expressions (LEEFTIJD, 23 cases) were the main
381 contributors, reflecting surface ambiguity with institution-like abbreviations
382 and clinical numeric values respectively.

383 Detection-only evaluation is incomplete without assessing whether residual
384 PHI remains recognizable in the final output. The reader study provides
385 empirical evidence for the practical protection offered by the complete pipeline.
386 After HIPS processing, readers collectively identified only five of 21 evaluable
387 leaked spans, with no individual reader identifying more than three (DS:
388 14.3%; 95% CI 3.0–36.3%). Simultaneously, readers generated 24–55 false
389 positives each, consistently flagging correctly replaced surrogates as suspicious
390 while overlooking genuine residual PHI.

391 The surrogate dilution analysis provides a mechanistic interpretation
392 of this finding. The 100 reviewed documents contained $S = 1,039$ HIPS
393 surrogate replacements alongside $L = 21$ evaluable leaked spans, giving a
394 theoretical precision ceiling of $L/(L + S) = 2.0\%$ for any reader who cannot
395 distinguish surrogates from genuine leaks by surface form alone. All observed

396 reader precisions (0%–8.3%) were consistent with random selection from this
397 pool under per-document Monte Carlo simulation. Protection is therefore not
398 solely a function of detection recall, as the large volume of realistic surrogates
399 structurally dilutes residual leaks to the point where they are statistically
400 indistinguishable from correctly replaced content.

401 This work also has limitations that point to important directions for
402 future research. External validation was limited to pathology, radiology, and
403 emergency care letter types from JBZ. Equivalent performance on outpatient
404 and clinical letters from unseen institutions cannot be assumed. Rare PHI
405 categories remain challenging despite augmentation, and the synthetic reports
406 introduce a dependency on the quality of the generative model. Finally, this
407 study does not provide formal re-identification risk quantification. Detection
408 recall and reader performance are proxies, and their relationship to actual
409 re-identification probability under a realistic adversary model remains an
410 open question.

411 Beyond the Dutch setting, the bootstrapping approach using a preliminary
412 model to enable GDPR-compliant cross-institutional data transfer through
413 automated tagging and surrogate replacement may offer a practical model
414 for other language communities assembling multicenter corpora under similar
415 regulatory constraints. The results also suggest that HIPS-based pipelines
416 warrant broader adoption. The residual risk profile of a well-calibrated
417 surrogate replacement system may be substantially more favorable than
418 detection recall figures alone imply, as even imperfect detection can provide
419 strong practical protection.

420 5. Conclusion

421 We presented GHMSCHRFT, an end-to-end de-identification pipeline
422 for Dutch clinical free-text combining a fine-tuned transformer model for
423 PHI detection with a Dutch-specific HIPS surrogate replacement module.
424 The pipeline achieved strong, consistent performance across internal and
425 external test sets, demonstrating cross-institutional generalizability for radiol-
426 ogy, pathology, and emergency care reports, and substantially outperforming
427 rule-based baselines. The reader study further showed that after surrogate
428 replacement, residual PHI was difficult for domain-familiar readers to distin-
429 guish from surrounding HIPS-generated content. Together, high detection
430 recall and realistic surrogate replacement make the pipeline robust to the
431 imperfect detection that is inevitable in practice. By releasing the detection

432 model, Dutch HIPS module, and annotation guidelines, this work provides
433 both a practical tool for Dutch clinical text de-identification and a reusable
434 framework for other clinical NLP settings operating under similar regulatory
435 constraints.

436 **Declaration of Competing Interest**

437 J.S. Bosma is employed by kaiko.ai. B. van Ginneken is a shareholder
438 of Thirona BV and Plain Medical. M. Prokop reports grants paid to his
439 institution from the Dutch Cancer Society (KWF), the European Commission,
440 and the Dutch government (ZonMW); honoraria and royalties paid to his
441 institution from Canon Medical Systems and Bracco SA; patents held by
442 Radboud University Medical Center, including one jointly with Canon Medical
443 Systems; an advisory board role at Fraunhofer MEVIS; leadership roles at
444 the European Society of Radiology and the Dutch Radiological Society; and
445 is a co-founder and shareholder of Plain Medical. All other authors declare
446 that they have no competing interests.

447 **Declaration of Generative AI and AI-Assisted Technologies in the** 448 **Writing Process**

449 During the preparation of this work, the authors used Claude (Anthropic)
450 for writing and coding assistance. After using this tool, the authors reviewed
451 and edited the content as needed and take full responsibility for the content
452 of this publication.

453 **Data Availability**

454 The clinical data used in this study are not publicly available due to
455 patient privacy constraints and applicable data protection regulations.

456 **Code Availability**

457 The code for running the GHMSCHRFT pipeline is publicly available
458 at <https://github.com/DIAGNijmegen/GHMSCHRFT>. The code for training
459 the PHI detection model and running all evaluation for this manuscript is
460 available at <https://github.com/DIAGNijmegen/GHMSCHRFT-training>.

461 CRediT authorship contribution statement

462 **Luc Builtjes:** Conceptualization, Methodology, Software, Formal anal-
463 ysis, Investigation, Data curation, Visualization, Writing – original draft,
464 Writing – review & editing. **Joeran S. Bosma:** Conceptualization, Method-
465 ology, Software, Data curation, Writing – review & editing. **Judith Lefkes:**
466 Validation, Data curation, Writing – review & editing. **Anouk Veldhuis:**
467 Data curation, Writing – review & editing. **Jeroen Geerdink:** Data cura-
468 tion, Writing – review & editing. **Tijmen de Haas:** Data curation, Writing –
469 review & editing. **Fabian R. Hoitsma:** Data curation. **Rianne A. Weber:**
470 Investigation, Writing – review & editing. **Sarah de Boer:** Investigation,
471 Writing – review & editing. **Myrthe Buser:** Investigation, Writing – review
472 & editing. **Harry J. Kendall:** Investigation, Writing – review & editing.
473 **Delano Sanches:** Investigation, Writing – review & editing. **Bram van**
474 **Ginneken:** Supervision, Writing – review & editing. **Henkjan Huisman:**
475 Supervision, Writing – review & editing. **Ewoud J. Smit:** Supervision,
476 Writing – review & editing. **Mathias Prokop:** Supervision, Writing – review
477 & editing. **Alessa Hering:** Conceptualization, Supervision, Writing – review
478 & editing.

479 References

- 480 [1] K. Kreimeyer, M. Foster, A. Pandey, N. Arya, G. Halford, S. F. Jones,
481 R. Forshee, M. Walderhaug, T. Botsis, Natural language processing
482 systems for capturing and standardizing unstructured clinical information:
483 a systematic review, *Journal of Biomedical Informatics* 73 (2017) 14–29.
484 doi:10.1016/j.jbi.2017.07.012.
- 485 [2] B. Percha, Modern clinical text mining: a guide and review, *Annual*
486 *Review of Biomedical Data Science* 4 (1) (2021) 165–187. doi:10.1146/
487 annurev-biodatasci-030421-030931.
- 488 [3] T. M. Seinen, E. A. Fridgeirsson, S. Ioannou, D. Jeannetot, L. H. John,
489 J. A. Kors, A. F. Markus, V. Pera, A. Rekkas, R. D. Williams, et al., Use
490 of unstructured text in prognostic clinical prediction models: a systematic
491 review, *Journal of the American Medical Informatics Association* 29 (7)
492 (2022) 1292–1302. doi:10.1093/jamia/ocac058.
- 493 [4] A. Rajkomar, E. Oren, K. Chen, A. M. Dai, N. Hajaj, M. Hardt, P. J. Liu,
494 X. Liu, J. Marcus, M. Sun, et al., Scalable and accurate deep learning

- 495 with electronic health records, *npj Digital Medicine* 1 (1) (2018) 18.
496 doi:10.1038/s41746-018-0029-1.
- 497 [5] European Union, Regulation (EU) 2016/679 of the European Parliament
498 and of the Council, [https://eur-lex.europa.eu/eli/reg/2016/679/](https://eur-lex.europa.eu/eli/reg/2016/679/oj)
499 oj, general Data Protection Regulation (2016).
- 500 [6] S. M. Meystre, F. J. Friedlin, B. R. South, S. Shen, M. H. Samore,
501 Automatic de-identification of textual documents in the electronic health
502 record: a review of recent research, *BMC Medical Research Methodology*
503 10 (1) (2010) 70. doi:10.1186/1471-2288-10-70.
- 504 [7] A. E. W. Johnson, L. Bulgarelli, T. J. Pollard, Deidentification of free-
505 text medical records using pre-trained bidirectional transformers, in:
506 Proceedings of the ACM Conference on Health, Inference, and Learning,
507 2020, pp. 214–221. doi:10.1145/3368555.3384455.
- 508 [8] A. Kovačević, B. Bašaragin, N. Milošević, G. Nenadić, De-identification
509 of clinical free text using natural language processing: a systematic
510 review of current approaches, *Artificial Intelligence in Medicine* 151
511 (2024) 102845. doi:10.1016/j.artmed.2024.102845.
- 512 [9] I. C. Wiest, M.-E. Leßmann, F. Wolf, D. Ferber, M. v. Treeck, J. Zhu,
513 M. P. Ebert, C. B. Westphalen, M. Wermke, J. N. Kather, Deidentifying
514 medical documents with local, privacy-preserving large language models:
515 the LLM-anonymizer, *NEJM AI* 2 (4) (2025) AIdbp2400537. doi:
516 10.1056/AIdbp2400537.
- 517 [10] Z. Liu, Y. Huang, X. Yu, L. Zhang, Z. Wu, C. Cao, H. Dai, L. Zhao,
518 Y. Li, P. Shu, et al., DeID-GPT: zero-shot medical text de-identification
519 by GPT-4, arXiv preprint arXiv:2303.11032 (2023). doi:10.48550/arXiv.2303.11032.
- 520
521 [11] K. Aghakasiri, N. Zambare, J. Thai, C. Ye, M. Mehta, J. R. Mitchell,
522 M. Abdalla, Not what the doctor ordered: surveying LLM-based de-
523 identification and quantifying clinical information loss, in: Proceedings
524 of the 2025 Conference on Empirical Methods in Natural Language
525 Processing, 2025, pp. 32175–32191. doi:10.18653/v1/2025.emnlp-main.1637.
526

- 527 [12] N. Carlini, F. Tramer, E. Wallace, M. Jagielski, A. Herbert-Voss, K. Lee,
528 A. Roberts, T. Brown, D. Song, U. Erlingsson, et al., Extracting training
529 data from large language models, in: 30th USENIX Security Symposium
530 (USENIX Security 21), 2021, pp. 2633–2650.
- 531 [13] OpenMed, PII and de-identification models, <https://huggingface.co/collections/OpenMed/pii-and-de-identification> (2023).
- 533 [14] OpenAI, OpenAI privacy filter, <https://huggingface.co/openai/privacy-filter> (2026).
- 535 [15] V. Menger, F. Scheepers, L. M. van Wijk, M. Spruit, DEDUCE: a pattern
536 matching method for automatic de-identification of Dutch medical text,
537 Telematics and Informatics 35 (4) (2018) 727–736. doi:10.1016/j.te
538 le.2017.08.002.
- 539 [16] D. Carrell, B. Malin, J. Aberdeen, S. Bayer, C. Clark, B. Wellner,
540 L. Hirschman, Hiding in plain sight: use of realistic surrogates to reduce
541 exposure of protected health information in clinical text, Journal of
542 the American Medical Informatics Association 20 (2) (2013) 342–348.
543 doi:10.1136/amiajnl-2012-001034.
- 544 [17] D. S. Carrell, B. A. Malin, D. J. Cronkite, J. S. Aberdeen, C. Clark,
545 M. Li, D. Bastakoty, S. Nyemba, L. Hirschman, Resilience of clinical
546 text de-identified with “hiding in plain sight” to hostile reidentification
547 attacks by human readers, Journal of the American Medical Informatics
548 Association 27 (9) (2020) 1374–1382. doi:10.1093/jamia/ocaa095.
- 549 [18] L. Builtjes, J. S. Bosma, A. Hering, DIAGNijmegen/dutch-med-hips:
550 v1.0.1-zenodo (2026). doi:10.5281/zenodo.18415297.
- 551 [19] H. Nakayama, T. Kubo, J. Kamura, Y. Taniguchi, X. Liang, doccano:
552 text annotation tool for humans (2018).
553 URL <https://github.com/doccano/doccano>
- 554 [20] J. Bosma, dragon-roberta-large-mixed-domain (revision 711c1d3) (2024).
555 doi:10.57967/hf/2170.
556 URL <https://huggingface.co/joeranbosma/dragon-roberta-large-mixed-domain>
557

- 558 [21] J. S. Bosma, K. Dercksen, L. Builtjes, R. André, C. Roest, S. J. Fransen,
559 C. R. Noordman, M. Navarro-Padilla, J. Lefkes, N. Alves, et al., The
560 DRAGON benchmark for clinical NLP, *npj Digital Medicine* 8 (1) (2025)
561 289. doi:10.1038/s41746-025-01626-x.
- 562 [22] A. Yang, A. Li, B. Yang, B. Zhang, B. Hui, B. Zheng, B. Yu, C. Gao,
563 C. Huang, C. Lv, et al., Qwen3 technical report, arXiv preprint
564 arXiv:2505.09388 (2025). doi:10.48550/arXiv.2505.09388.
- 565 [23] OpenMed, OpenMed/privacy-filter-nemotron: fine-grained pii extraction
566 with 55 categories, <https://huggingface.co/OpenMed/privacy-filter-nemotron>
567 (2026).